

MAXIMIZING UTILITY OF GENOME SEQUENCE DATA

David J. Dooling, Lynn Carmichael, Li Ding, Craig S. Pohl,
Scott Smith, Joshua McMichael, George M. Weinstock,
Elaine R. Mardis, and Richard K. Wilson
The Genome Center at Washington University
in St. Louis School of Medicine

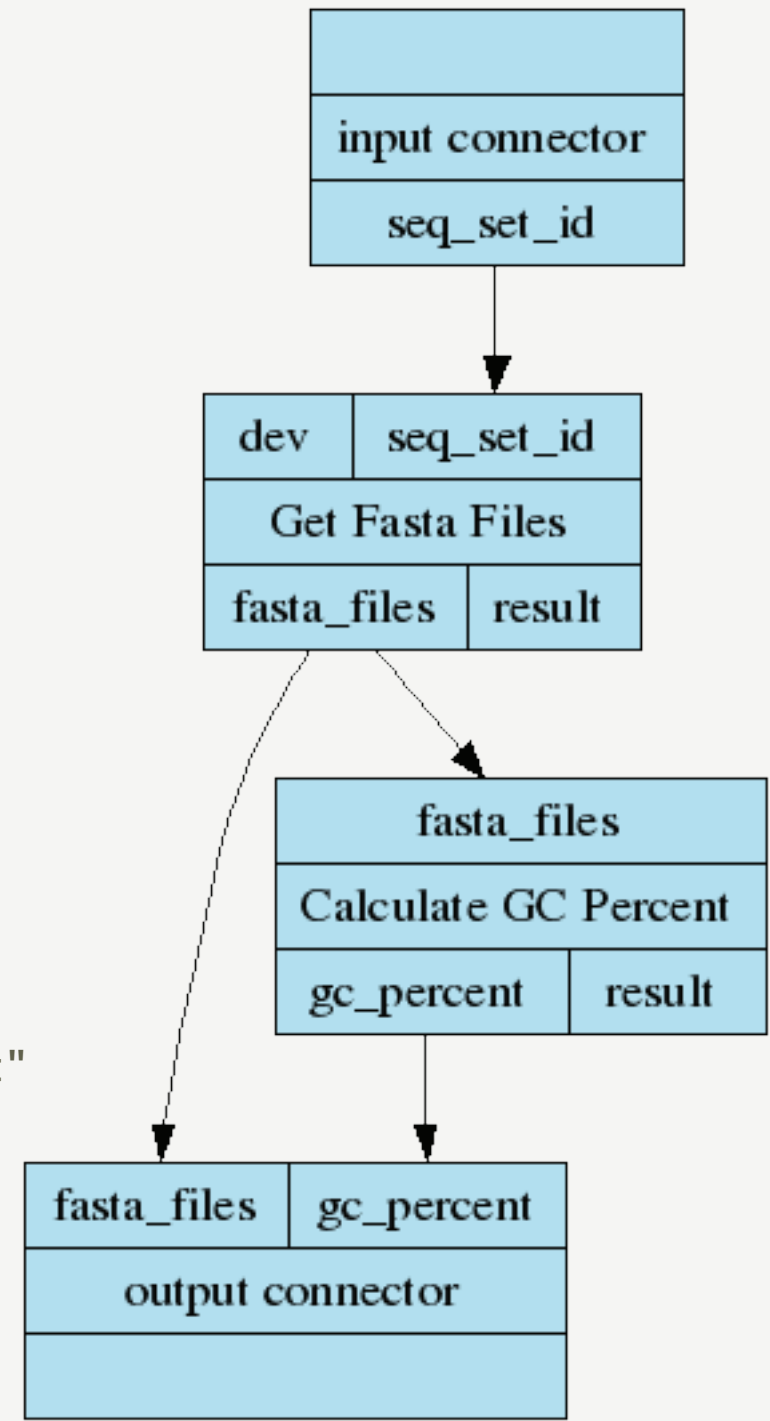
ABSTRACT

Advances in DNA sequencing technologies over the past few years have led to increases in data generation and processing rates that far outpace Moore's Law and storage capacity improvements. As a result, there will come a time when one will no longer be able to “throw more money” at the problems presented by DNA sequencing, i.e., researchers will not be able to keep pace with data generation by purchasing more and more storage and computational nodes. Proposed sequencing platform improvements and the rapid rate of adoption of these technologies by labs large and small will only hasten the time when the old solutions will no longer apply. The history of freely shared sequence data through the NCBI and EBI Trace Archives transform the very difficult problem of massive sequence data generation into a problem of data generation and data sharing on a scale heretofore unimaginable. Over the last year, several organizations, e.g., MGED, NCI, Illumina, 1000 Genomes DCC, and NHGRI, have convened meetings to discuss the problems presented by the massive amounts of data generated by next-generation sequencing technologies. As prologue, brief overviews of these meetings will be presented along with approaches to dealing with massive data generation rates from other disciplines, e.g., high energy physics and high-resolution medical imaging. The Genome Center at Washington University in St. Louis, due to its large-scale sequencing operation and whole-genome analysis capabilities, experiences the difficulties presented by massively-parallel sequencing platforms acutely. To address the many challenges presented by the scale of data generation and requisite analysis, we have developed a multidisciplinary approach involving experts in biology, genomics, bioinformatics, computer science, information technology, and engineering. The resulting approach involves many techniques including intelligent compression and data reduction, data aging, archiving, parallelization, fault-tolerant workflows, scalable software frameworks, and multivariate / multi-genome visualization and comparison, which leverage and extend our laboratory information management system. This approach and its application to the sequencing and analysis of cancer samples will be presented.

UR

- Object Relational Mapping (ORM) layer
- Distributed, in-memory transactions (STM)
- Manage object state between application and persistence layer (RDBMS)
- Highly parallel and high throughput
- Written in Perl

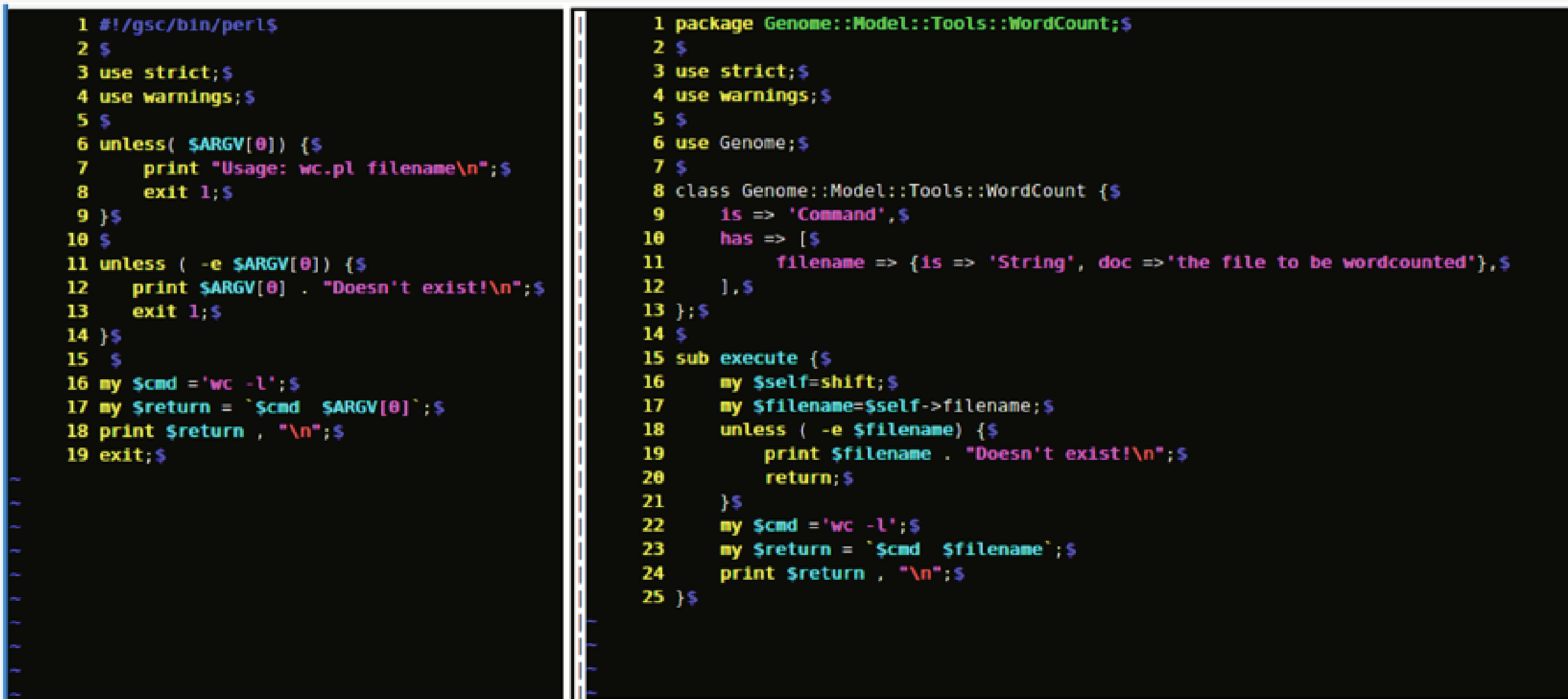
Genome :: Workflow



Genome :: Model

- Represents a state of belief about a genome
- Genome model information includes:
 - Patient/sample phenotype
 - Sequence, array, etc. data
 - Software versions and configuration
 - Analysis results
- A single sample/patient/population could have several concurrent models
- Built on UR and Genome::Workflow

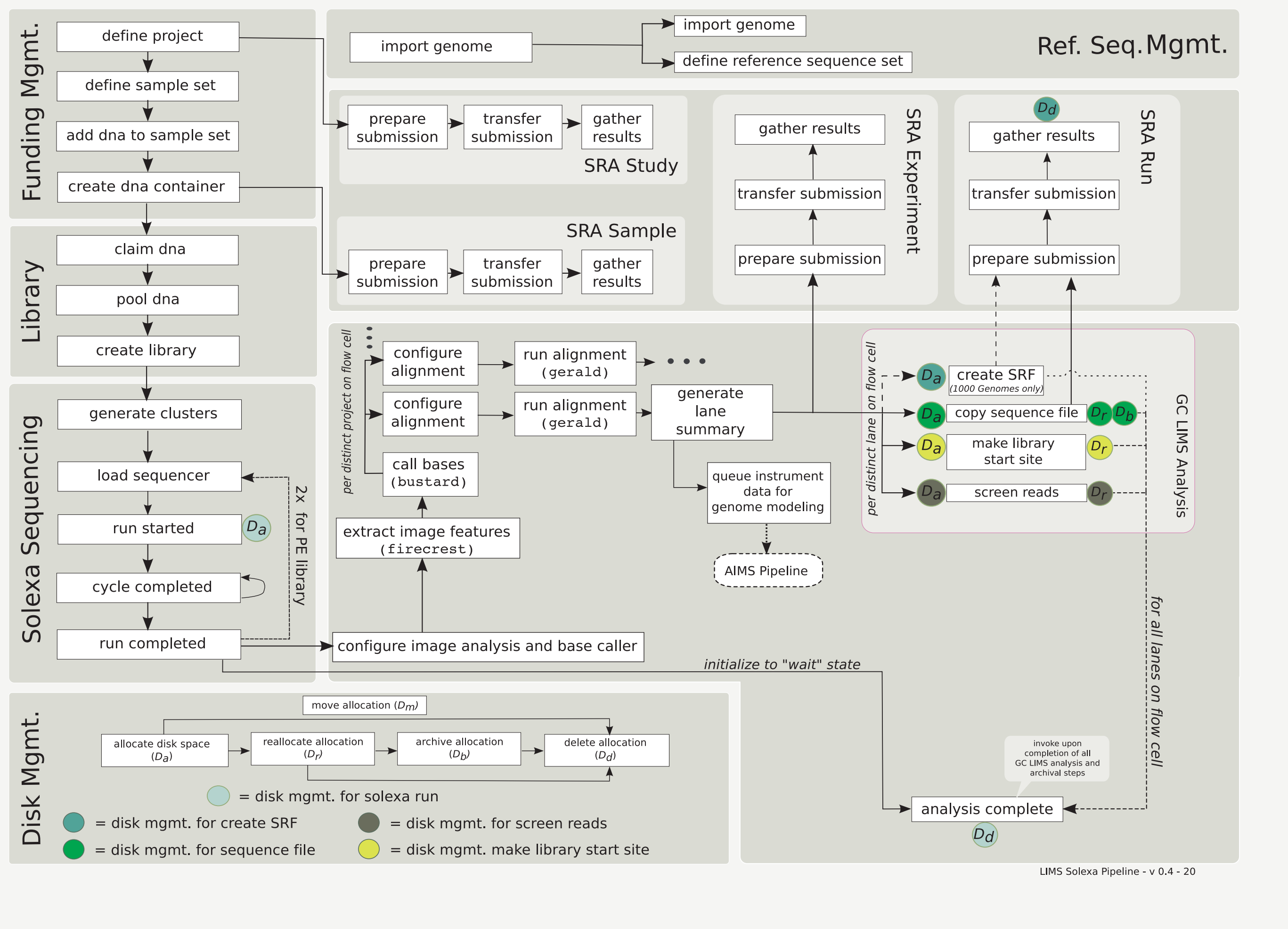
Genome :: Model :: Tools



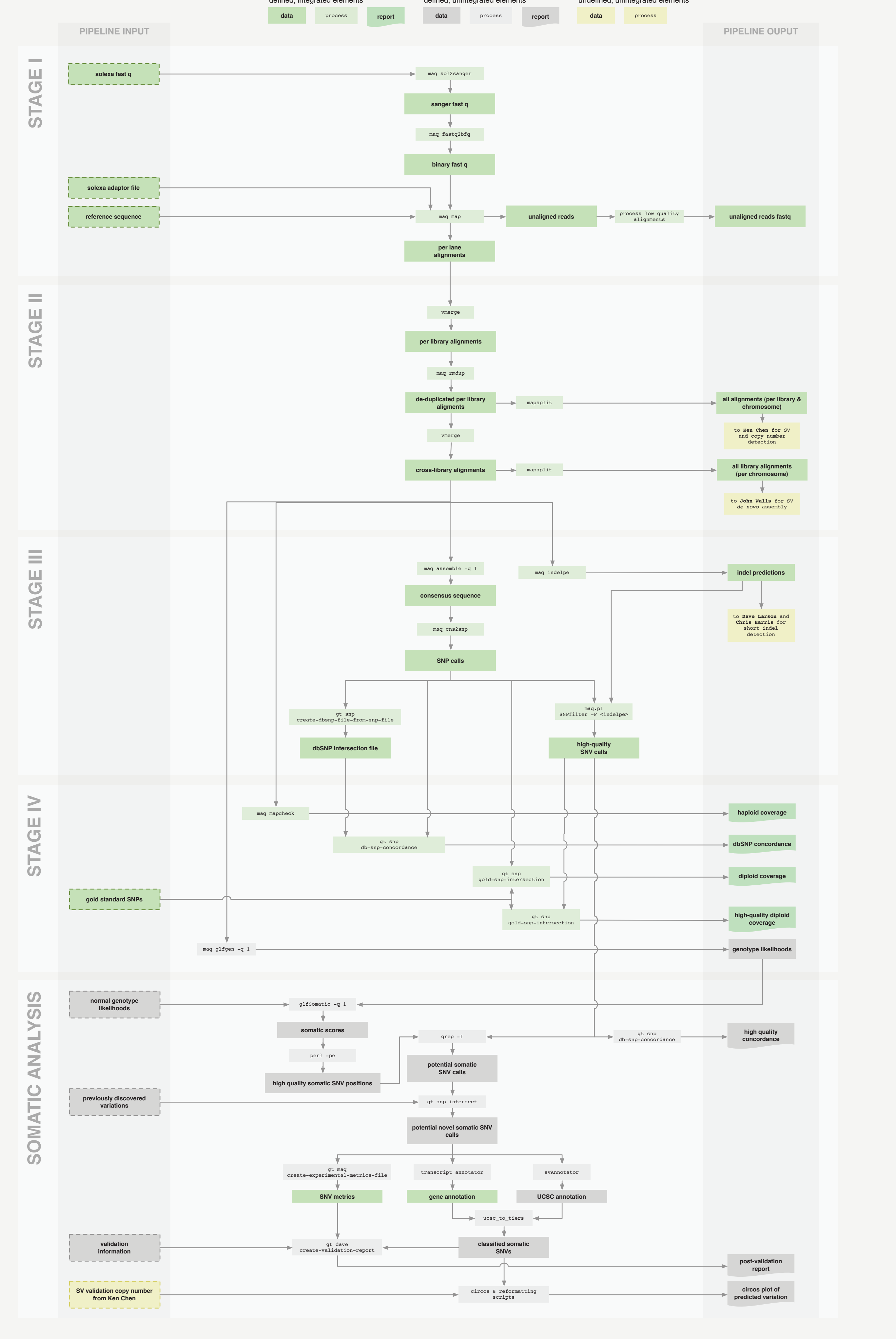
Genome :: Model :: Reports



LIMS SOLEXA PIPELINE



VARIATION DETECTION PIPELINE



FUTURE WORK

- Public release (GPL3)
- Improved reports
 - Pipeline status and result reports
 - Enhanced genome-model comparison tools
- High-throughput computing

